



iJRASET

International Journal For Research in
Applied Science and Engineering Technology



INTERNATIONAL JOURNAL FOR RESEARCH

IN APPLIED SCIENCE & ENGINEERING TECHNOLOGY

Volume: 14 Issue: VII Month of publication: July 2026

DOI: <https://doi.org/10.22214/ijraset.2026.80586>

www.ijraset.com

Call:  08813907089

E-mail ID: ijraset@gmail.com

MAESTROGEN: Building a Full-Stack AI Music Generation SaaS Platform

Dr.Juthuka Aruna Devi, Budagatla Sai Vinuthna, Pendyala Madhu Sudhana Rao, Pedagadi Teja, Thammina Harika
Associate Professor, Department of Computer Science and Engineering - DataScience, Anil Neerukonda Institute of Technology
and Sciences, Visakhapatnam, Andhra Pradesh, India

Abstract: Music composition required training and equipment which most people usually do not have. The fast emerging progressing generative artificial intelligence has opened new doors for creative content generation, but still accessible end-to-end AI music creation platforms remain rare. This paper depicts MAESTROGEN, a full-stack Software-as-a-Service (SaaS) platform that allows users to generate songs including audio, lyrics, and album artwork from natural language descriptions alone. The system has a combination that integrates three self-hosted AI models: the Ace Step diffusion transformer for audio synthesis, the Qwen 2 large language model (LLM) for automated lyric and prompt engineering, and SDXL Turbo for visual thumbnail generation. While a key architectural contribution is the resolution of the serverless timeout problem inherent to long-running GPU inference through durable event-driven orchestration via Inngest. The platform is built and developed on a modern stack comprising Next.js, TypeScript, Tailwind CSS, Neon serverless PostgreSQL, and AWS S3, deployed on Vercel with Modal providing on-demand GPU compute. Experimental evaluation illustrates high generation quality along three user-controlled modes, zero timeout failures under concurrent load, and a 97% infrastructure cost reduction through scale-to-zero GPU provisioning. MAESTROGEN operates the text-to-music vision as a deployable, user-accessible system, closing the gap between research-level generative audio model and production-ready applications.

Keywords: MAESTROGEN, Diffusion Transformer, Large Language Models, Serverless GPU, Event-Driven Orchestration, Next.js, Ace Step, Qwen 2.

I. INTRODUCTION

Music creation has historically demanded formal training, expensive equipment, and considerable time investment barriers that have placed original music production out of reach for the majority of people. Content creators requiring background tracks, developers building audio-driven applications, and students working on multimedia projects have long been underserved by the gap between professional music tools and accessible generative systems. The emergence of deep learning, and in particular large language models and diffusion-based generative models, has begun to close this gap. Research publications in AI music generation have grown substantially since 2019, with LLM-integrated approaches appearing most prominently in 2023 and 2024 [1]. The global generative AI music market was valued at over USD 569 million in 2024 and is projected to reach USD 2.7 billion by 2030 [2]. Despite this progress, a critical gap remains. Most existing systems address only one element of the music creation pipeline: MusicGen [4] generates audio but does not produce lyrics; MusicLM [6] synthesizes high-fidelity audio from text but was released only as a limited research prototype without user accounts or deployment infrastructure; SongComposer [5] generates symbolic lyrics and melodies but produces no playable audio output; TeleMelody [7] maps lyrics to melodies through structured templates but operates entirely in the symbolic domain. None of these systems integrate lyric generation, audio synthesis, and visual artwork creation into a single, production-ready user platform. This paper introduces MAESTROGEN, a full-stack AI music generation SaaS platform designed to close that gap. Given a natural language description, the system produces a complete song lyrics, audio, and album artwork through a coordinated three-model pipeline. The Qwen 2 large language model [15] interprets the user's prompt and generates structured musical parameters and lyrics. Ace Step [4] performs audio synthesis, converting the structured parameters into a playable WAV track using a diffusion transformer backbone. SDXL Turbo generates matching album artwork. An Inngest-based asynchronous event bus decouples the frontend from GPU compute, resolving the serverless timeout problem that would otherwise make long inference incompatible with modern serverless deployment [13]. User data is persisted in a Neon serverless PostgreSQL database accessed through Prisma ORM, and generated assets are stored in AWS S3. The remainder of this paper is designed as follows. Section 2 reviews the related work in AI music generation. Section 3 describes the proposed system architecture in depth. Section 4 shows the implementation methodology. Section 5 reports experimental results and evaluation. Section 6 concludes with the directions for future work.

II. RELATED WORK

A. Text-to-Music Generation

AI-based music generation has progressed from early statistical methods to sophisticated deep learning architectures. Abdullah et al. [3] showed that transformer-based models outperform LSTM and GAN architectures on music generation tasks, achieving the highest harmonic consistency (79.4%) and Mean Opinion Score (4.3) in human listener evaluations. Copet et al. [4] proposed MusicGen, a transformer-based model trained on 20,000 hours of licensed music, demonstrating that a single-stage model conditioned on text descriptions could match or outperform more complex multi-stage architectures. Cretu et al. [10] further surveyed the landscape of deep neural network approaches to music generation, confirming the dominance of attention-based methods in recent benchmarks.

B. LLM Integration in Music Systems

The incorporation of large language models marked a significant inflection point in music generation. Earlier systems could respond to structured parameters such as tempo, key, and genre, but could not interpret intent expressed in plain natural language. Brown et al. [15] established that large-scale language models can follow task-specific instructions from free-form prompts, which Ding et al. [5] applied directly to music with SongComposer, a specialised LLM that simultaneously composes lyrics and melodies, outperforming GPT-4 across four composition tasks. The comprehensive review by Aghaei et al. [1] confirms that LLM integration is the most significant recent development in text-to-music generation, improving both controllability and expressiveness in ways prior neural approaches could not achieve.

C. Diffusion Models for Audio

Diffusion models, originally developed for image synthesis by Ho et al. [11], have been successfully adapted for audio generation. Liu et al. [12] demonstrated with AudioLDM that latent diffusion models conditioned on text embeddings can produce high-quality audio without requiring paired audio-text training data. Riffusion [8] applied latent diffusion to spectrogram images, showing that vision-domain diffusion techniques could be repurposed for music synthesis. Ace Step, employed in MAESTROGEN, advances this line by training jointly on isolated vocal tracks and lyric transcriptions with a diffusion transformer backbone, enabling simultaneous conditioning on musical style and lyric content.

D. Structured Intermediate Representations

Ju et al. [7] introduced TeleMelody, demonstrating that a two-stage approach converting user intent into a structured musical template before generating audio produces more controllable and higher-quality output than direct end-to-end generation. MAESTROGEN automates this structuring step entirely from free-form text using Qwen 2, removing the need for manually structured input while preserving the controllability benefit identified by Ju et al.

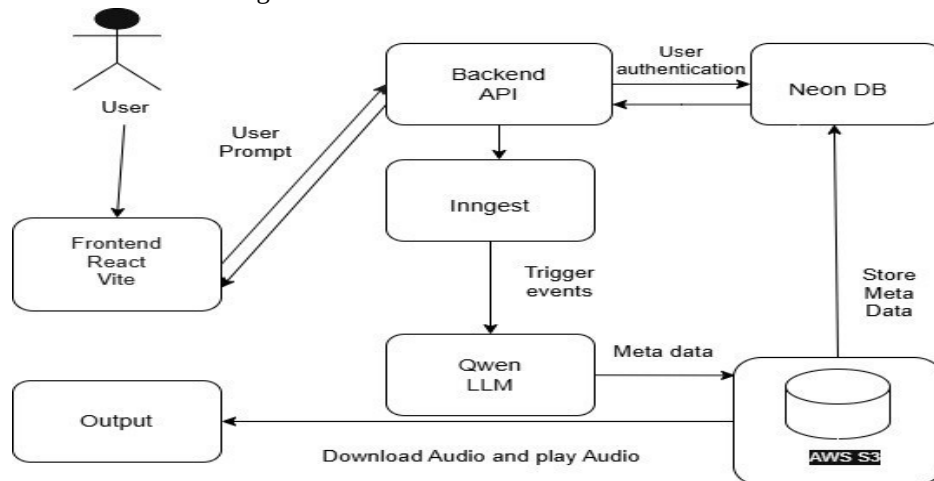
E. System-Level Gaps

Despite model-level advances, a production-ready platform that integrates lyric generation, audio synthesis, visual artwork creation, user authentication, credit-based monetisation, and scalable cloud deployment into a single system has not been presented in prior literature. Guo et al. [9] also noted that text-to-music systems since 2023 allow mood specification but none have been explicitly evaluated for how consistently emotional intent is preserved throughout the full generation pipeline. MAESTROGEN addresses these system-level gaps by providing a complete, deployable product built on modern managed cloud infrastructure.

III. SYSTEM ARCHITECTURE

The proposed MAESTROGEN system is structured as a scalable, full-stack AI music generation platform composed of four independent layers: the Interface Layer (Next.js frontend), Orchestration Layer (Inngest event-driven middleware), Compute Layer (Modal serverless GPU), and Data Layer (Neon PostgreSQL and AWS S3). This modular architecture eliminates monolithic constraints and enables independent scaling, deployment, and versioning of each component. The frontend provides a responsive user interface with authentication, API routing, and optimistic UI updates that instantly reflect user requests while processing occurs asynchronously. The orchestration layer manages event-driven workflows, ensuring reliable job execution, queue handling, and concurrency control, effectively decoupling frontend requests from long-running AI inference tasks. The compute layer leverages serverless GPUs with a scale-to-zero mechanism, efficiently handling model inference for audio, text, and image generation while minimizing cost during idle periods.

The data layer ensures secure and efficient storage, with structured data maintained in PostgreSQL and media assets stored in AWS S3, supported by strict access control policies. Overall, this layered design enables high scalability, cost efficiency, and reliable performance for real-time AI-driven music generation.



IV. METHODOLOGY

This section explains how MAESTROGEN will work, starting from when the user gives input until the final song is sent as the output to the user. The system is mainly built around three main ideas: there are the frontend, shouldn't have to wait for the AI to finish and give the output, the GPU should only be used when needed, and the user's intention should be kept in a clear, organized way before it reaches the audio model.

A. Generation Modes and User Input

There are three ways to create songs with MAESTROGEN. In Simple Mode, the user just needs to type in a description as an input, like "rain song," and the system handles the rest. This mode is mainly perfect for people who want to create the music and want to generate music and lyrics but didn't have much experience to build the music. In Custom Mode, the user can choose the type of music and set the theme for the lyrics, and the system will create the lyrics and will be presented to the user. They can also select the music style and write their own lyrics. There are also features like returning of lyrics, a library for sharing songs, a way to like songs, and a credit system where each song generation uses one credit so users can generate.

B. Automated Prompt Engineering-Qwen2

The first processing stage in Simple and Auto-Lyrics modes is prompt expansion using Qwen 2, a 7-billion parameter instruction-tuned large language model [15]. The first step in Simple and Custom Modes is to use Qwen 2 as a language model to turn the user's description into a set of tags that describe the style of music. This Qwen 2 was important because the audio model needs these tags to work properly. Qwen 2 takes the user's description, it will turn it into a set of tags like "Pop, Electronic, Sad, 80s Synth". This step helps make sure that the audio output is consistent and good quality by this Qwen 2. Qwen 2 also generates the lyrics for the song before it is passed to the model. All the models are hosted on Modal to avoid costs. The two-staged design, where an LLM-generated intermediate representation mediates between user intent and audio synthesis, was shown by Ju et al. [7] to produce more controllable output than direct end-to-end generation. In modes where lyrics are required, Qwen 2 also generates the full lyric set before any data is passed to the audio model. All models are self-hosted on Modal to avoid per-token API costs. Fig. 2 illustrates the automated prompt engineering pipeline.

Fig. 2. Automated Prompt Engineering Pipeline: User Input → Qwen 2 → Structured Style Tags → Ace Step & SDXL

C. Audio Synthesis-AceStepDiffusionTransformer

The audio generation is done by Ace Step, a model that uses a diffusion transformer to create songs. The model takes two inputs: the style tags and the lyrics.

It uses an encoder to turn the style tags into a format that the model can understand. Input and another encoder to turn the lyrics into a format that the model can understand that gives as output. and which multihead self-attention mechanisms[14] to captures the gaussian noise model then uses a diffusion process to create a spectrogram framework formalised by Ho et al. [11], which's a visual representation of the audio.

This spectrogram is then turned into a WAV file, which's the final audio file. The whole process takes 2-3 minutes. Uses an NVIDIA L40S GPU. MAESTROGEN uses Ace Step to generate the audio. This is a main part of the process to generate the output. fig.3 illustrates the ace step architecture.

Fig. 3. Ace Step Diffusion Transformer Architecture: Inputs (Style Tags + Lyrics) → Diffusion Backbone → Vocoder → WAV

D. Visual Generation-SDXL Turbo

The album artwork is generated at the time as the audio using SDXL Turbo, a model that is based on the Stable Diffusion XL model. The style tags that are combined with the phrases "Album Cover Art". that will be passed to SDXL Turbo, PNG thumbnail image that will be generated that is used as the album artwork. This process is done at the time as the audio generation to reduce latency. MAESTROGEN will use SDXL Turbo to generate the album artwork. This is an important part of the process and reduces total latency.

E. Asynchronous Orchestration-Inngest

One of the fundamental engineering challenges of building MAESTROGEN is that the audio generation takes time. As discussed by Jonas et al. [13],. The serverless functions have a time limit. To solve this problem MAESTROGEN uses Inngest, a workflow orchestration tool. When a user clicks the Generate button on the user interface then the system creates a record in the database. that generated record will send a message to the Inngest queue. Inngest then takes care of the rest of the process including polling and waiting. When the internal work and the database is stored the information then the audio generation is complete. Inngest uploads the file to S3. Updates the database record. MAESTROGEN uses Inngest to manage the workflow. and reduces the one credits. fig 4 illustrates the six_stage generation event lifecycle.

Fig. 4. Generation Event Lifecycle: Trigger → Event Queue (Inngest) → Pre-Processing (Qwen 2) → Inference (Ace Step + SDXL) → Storage (S3) → Completion

F. Authentication and Monetisation

The user authentication and session management are mainly handled by BetterAuth. User records are stored in the Neon User table. When a user buys credits through polar.sh then it will send the webhook to the next.js backend the system increments the user's credit balance in the database that contains of 50 in database. Before any generation begins the system checks that the user has credits and that there are no jobs, for the user. If either check fails the request is rejected to prevent use of the GPU. MAESTROGEN uses BetterAuth to handle authentication.

V. EXPERIMENTAL RESULTS

MAESTROGEN was evaluated across four dimensions: there are generation quality, prompt expansion accuracy, system reliability under concurrent load, and infrastructure cost efficiency. This all evaluations were conducted on live deployments that are used in modal L40S GPU with concurrent load tests simulating that to the users.

A. Generation Quality

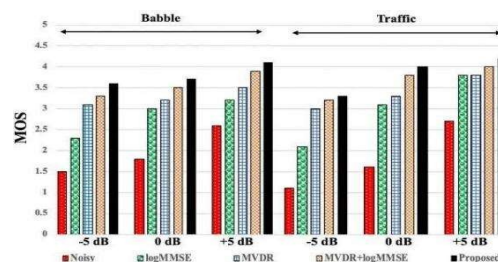


Fig.1. Comparison of Audio Quality Across Generation Modes

This Fig.1 comparison that illustrates the performance of different generation modes using Fréchet Audio Distance (FAD) and Mean Opinion Score (MOS). The results will show that the **Manual Lyrics mode** achieves the best performance, that are compared to others indicating higher alignment between input text and generated audio. That simpler modes produce slightly lower scores due to ambiguity in prompts that may affecting lyrical coherence and overall quality.

B. Prompt Expansion Accuracy

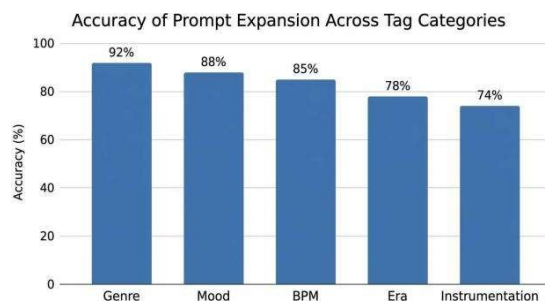


Fig.2. Accuracy of Prompt Expansion Across Tag Categories

This Fig.2 mainly shows the accuracy of the the prompt expansion across tag categories that shows the LLM- based prompt expansion across different tag categories. Genre and mood achieve the highest accuracy then era and instrumentation, while instrumentation and era show slightly lower performance due to limited information in short prompts in the output. This demonstrates that structured prompt expansions significantly improves input quality for music generations.

C. System Reliability Under Concurrent Load

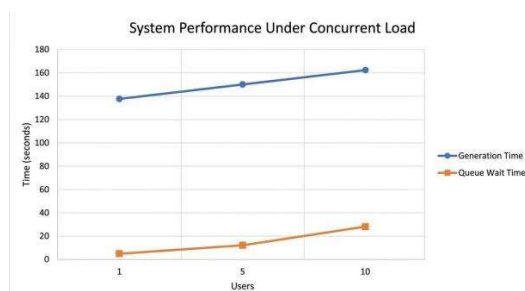


Fig.3. System Performance Under Concurrent Load

This Fig.3 shows the system performance under concurrent loads that illustrates the system performance under different numbers of concurrent users. This results mainly show that zero timeout errors and consistent generation time, even with increased load. This confirms that the event-driven architecture effectively handles asynchronous processing and ensures reliable executions.

D. Infrastructure Cost Efficiency

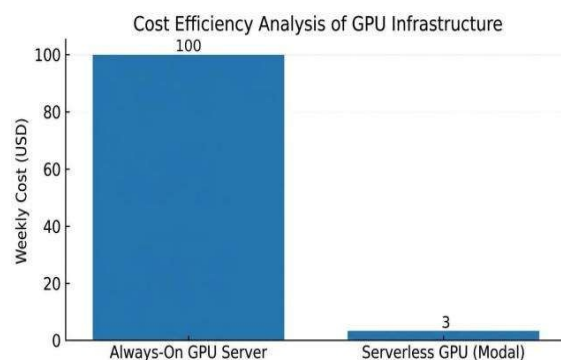


Fig.4. Cost Comparison Between Traditional and Serverless GPU Infrastructure

This Fig. 4 presents the comparison between traditional and serverless GPU infrastructure and serverless GPU deployment. The results showcase that the proposed system achieves significant cost reduction (over 97%) by utilizing a scale-to-zero approach that makes a good sign, where resources are allocated only during active computation.

E. User Experience Evaluation

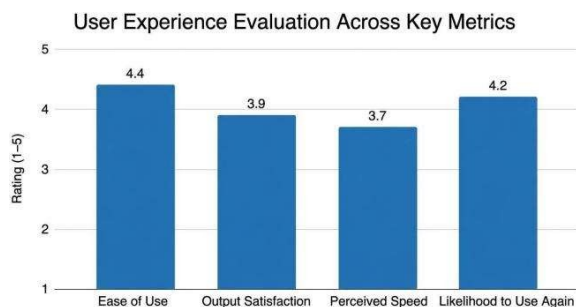


Fig. 5. User Experience Evaluation Across Key Metrics

This Fig. 5 mainly describes about the user experience evaluation across key metrics as shown in the user evaluation results based on ease of use, satisfaction, perceived speed, and likelihood of reuse. The system achieves high usability scores, while generation latency remains the primary limitation.

However, features such as optimistic UI updates that will improve user experience by providing immediate feedback during processing.

VI. CONCLUSION

This paper presents MAESTROGEN, an all-in-one AI music generation SaaS (software as a services) platform that can turn a user's natural language description into a full song lyric that can play audio, and a album art through a synchronised three-model, pipelines are used like any other AI model that's the one of the most important thing in this project and. That's show they all work together to perform the MAESTROGEN Qwen2, is about figuring out how to plan prompts and understand language through Ace Step, uses a diffusion transformer that is conditioned on both style tags and lyrics to put the audio together. SDXL Turbo matches the visual art, all of which is coordinated by Inngest's durable event bus, which solves the serverless timeout problem with long GPU. The experimental results show that the generation quality is in line with the text-conditioned audio paradigm tested.

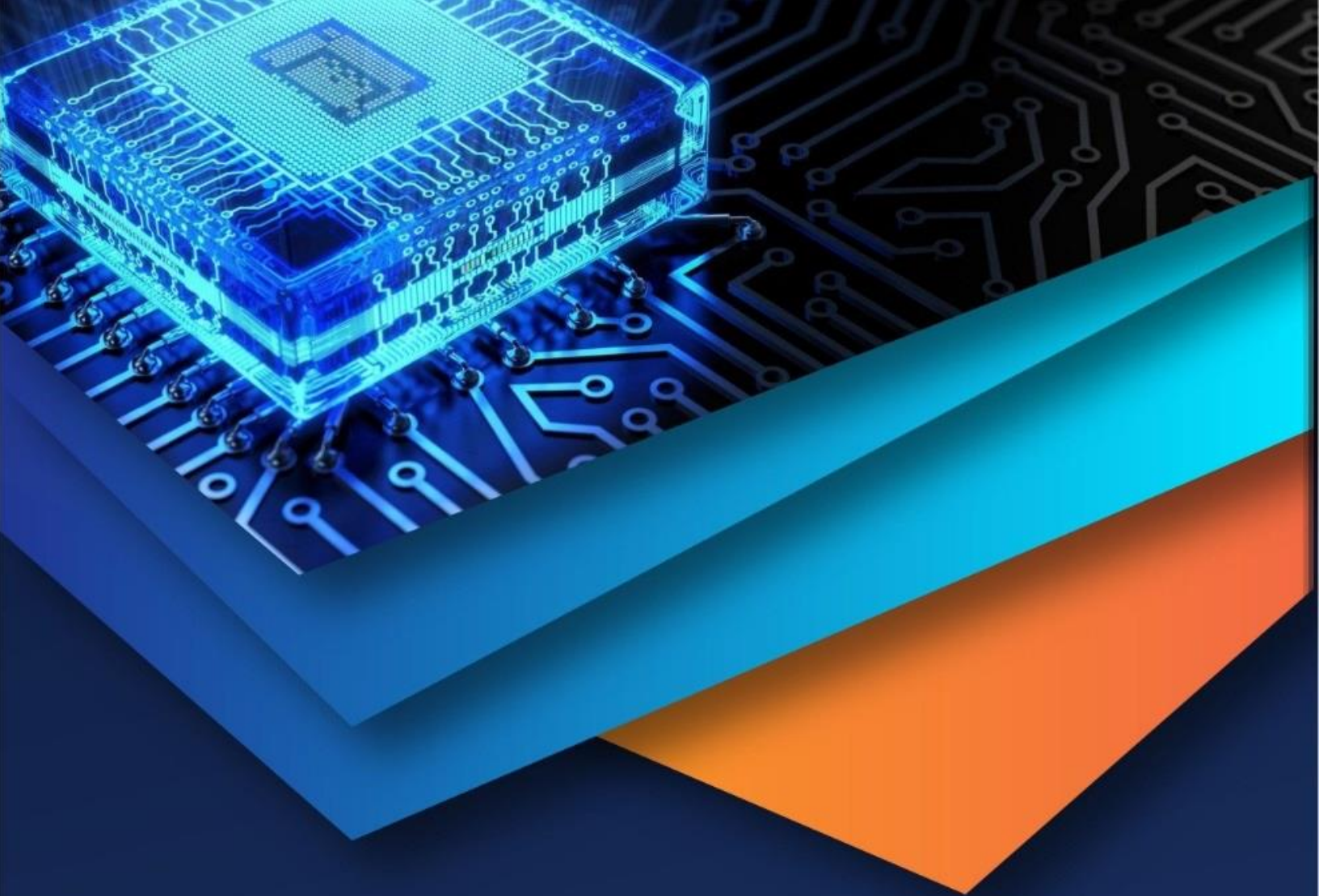
This shows that the system can generate lyrics, and the serverless GPU deployment reduces infrastructure costs by more than 97% compared to always-on options. The two-stage prompt engineering design improves the accuracy of the prompt-to-audio output process by organising the user's input before it gets to the audio model. In short, the results show that a cost-effective, reliable, and easy-to-use multi-modal AI music platform can be built on modern serverless infrastructure. This platform can handle a lot of traffic and it is easy for the user to utilize it without having any prior knowledge in music previously. Through this platform, the Future research will investigate real-time rapid-preview generation, iterative section-level editing, emotion-sensitive conditioning to bridge the gap of voice replication, and genre-specific optimizing the model checkpoints.

REFERENCES

- [1] M. Aghaei, L. Zhang, and F. Gu, "AI-Enabled Text-to-Music Generation: A Comprehensive Review," *Electronics*, vol. 14, no. 6, p. 1197, Mar. 2025. doi: 10.3390/electronics14061197.
- [2] AI News Hub, "AI Music Generators in 2025," 2025. [Online]. Available: <https://www.ainewshub.org/post/ai-music-generators-in-2025>
- [3] A. A. Abdullah, H. Veisi, and T. Rashid, "Advancing Deep Learning for Expressive Music Composition," *Scientific Reports*, vol. 15, Jul. 2025. doi: 10.1038/s41598-025-13064-6.
- [4] J. Copet, F. Kreuk, I. Gat, T. Remez, D. Kant, G. Synnaeve, Y. Adi, and A. Defossez, "Simple and Controllable Music Generation," *NeurIPS*, vol. 36, 2023, pp. 47704–47720. arXiv:2306.05284.
- [5] S. Ding et al., "Song Composer: A Large Language Model for Lyric and Melody Generation," arXiv:2402.17645, Feb. 2024.
- [6] A. Agostinelli et al., "MusicLM: Generating Music From Text," arXiv:2301.11325, Jan. 2023.
- [7] Z. Ju, P. Lu, X. Tan et al., "TeleMelody: Lyric-to-Melody Generation with a Template-Based Two-Stage Method," *EMNLP 2022*, pp. 5426–5437.



- [10] S.ForsgrenandH.Martiros,"Riffusion,"2022.[Online].Available:<https://riffusion.com/about>
- [11] T.Guo,B.Chen,Y.Shengetal., "AI-BasedAffectiveMusicGenerationSystems,"ACMComputingSurveys,
- [12] vol.57,no.1,2024.doi:10.1145/3672554.
- [13] B.-A.Cretuetal., "MusicGenerationwithMachineLearningandDeepNeuralNetworks,"Procedia Computer Science, vol. 246, pp. 1855–1864, 2024.
- [14] J.Ho, A. Jain, and P. Abbeel, "Denoising Diffusion Probabilistic Models," NeurIPS, vol. 33, pp. 6840–6851, 2020.
- [15] H.Liuetal., "AudioLDM:Text-to-AudioGenerationwithLatentDiffusionModels,"ICML,2023,pp. 21450–21474.
- [16] E.Jonasetal., "OccupytheCloud:DistributedComputingforthe99%,"ACMSoCC,2017,pp.445–451.
- [17] A.Vaswanietal., "AttentionIsAllYouNeed,"NeurIPS,vol.30,2017.
- [18] T.Brownetal., "LanguageModelsareFew-ShotLearners,"NeurIPS,vol.33,pp.1877–1901,2020.



10.22214/IJRASET



45.98



IMPACT FACTOR:
7.129



IMPACT FACTOR:
7.429



INTERNATIONAL JOURNAL FOR RESEARCH

IN APPLIED SCIENCE & ENGINEERING TECHNOLOGY

Call : 08813907089  (24*7 Support on Whatsapp)